

Building Your Ecosystem of AI Agents

Principles, Architecture, and Practice: A Teaching Guide

Preface: Why This Guide Exists

Something has changed in how we work with artificial intelligence. For the past few years, most of us have used AI the way we use a search engine or a calculator: we ask, it answers, and the exchange ends there. What is emerging now is different in kind, not merely in degree. We can now delegate. We can describe a responsibility, entrust it to a software agent, and expect that agent to carry the work forward over hours, days, and months, consulting us at the moments that matter and keeping a faithful record of what it has done.

This guide teaches you how to build not just one agent but an ecosystem of them: a small, well-ordered team in which each member has a clear purpose, the members share a common record, their conduct is governed by explicit principles, and a human being remains at the center of every consequential decision. It is written for three audiences at once. If you have never written a line of code, the early sections and the first of the six steps are for you, and you can build something real this week. If you work in health care or another field where trust and confidentiality are not optional, the sections on private agents, data boundaries, and safeguards speak directly to your situation. If you are a practitioner who wants full technical depth, the architecture chapters include the protocols, configurations, and repository structures you need.

One conviction runs through the whole document: the technology is the easy part. The harder and more valuable work is deciding what these agents are for, what principles govern their conduct, and how their efforts are harmonized so that the whole serves a purpose greater than any part. A collection of powerful tools without a unifying vision produces noise. A modest set of tools, arranged around clear principles and a shared record, produces something closer to an institution: small, but trustworthy.

Part One: Foundations

Chapter 1: What an Agent Is, and What It Is Not

An AI agent is a language model given three additional things: tools, memory, and a mandate.

The language model itself is the mind of the operation. On its own, a model can only converse: you give it text, it returns text. An agent wraps that mind in a body. Tools are its hands: the ability to search the web, read and write files, send an email, query a database, run a program. Memory is its record book: files, notes, and histories that persist after the conversation ends. The mandate is its character and commission: standing instructions that tell it who it serves, what it values, what it may do freely, and what it must never do without asking.

The most useful early distinction, articulated well in Anthropic's essay "Building Effective Agents," is between a workflow and an agent. In a workflow, you have decided the steps in advance: first summarize the document, then extract the dates, then file the result. The model fills in the steps, but the path is yours. In an agent, you specify the destination and the agent chooses the path: it decides which tool to use next, observes the result, and decides again, looping until the work is done or it needs your judgment. Workflows are more predictable and cheaper; agents are more capable and more surprising. A wise builder uses the simplest form that serves the purpose, and reaches for full agency only when the task requires it. Much of what people call "agents" in marketing is, usefully and honestly, workflows, and there is no shame in that.

A few units of measurement will serve you throughout this guide, and they are worth understanding rather than memorizing.

Tokens are the syllables of a language model. A token is roughly three quarters of an English word; this sentence is about twenty tokens. Every model reads and writes in tokens, and every commercial model is priced in them, typically in dollars per million. When you see a price like three dollars per million input tokens and fifteen dollars per million output tokens, translate it: a million tokens is roughly seven hundred fifty thousand words, about ten novels. Reading is cheap; writing is several times dearer; and agents, because they loop, read the same growing record again and again. This is why an agent task can quietly cost ten or a hundred times what a single question costs, and why Anthropic reported that its multi-agent research system used roughly fifteen times the tokens of an ordinary chat. Token awareness is not miserliness. It is the same discipline as knowing what a lab test costs before ordering one on every patient.

The context window is the desk, not the library. It is the amount of text a model can hold in view at one moment, now commonly two hundred thousand to a million

tokens. Whatever is not on the desk does not exist for the model, no matter how important it was yesterday. This single fact explains most of agent architecture: the entire craft of memory files, shared records, and summarization exists because the desk is finite and is cleared at the end of every session. When this guide later insists on a written record that agents consult, it is working with the grain of this limitation, not around it.

Latency and cost per task matter more than raw speed. A local model that answers in a tenth of a second and a frontier model that takes twenty seconds are not competitors; they are different instruments. You will learn to route small, private, repetitive work to the small fast instrument and reserve the large one for judgment.

One more clarification, because the word causes confusion: "agent" describes a role, not a product. The same underlying model can serve as your legal reviewer in the morning and your research curator in the afternoon. What distinguishes the agents in your ecosystem is not which company's model animates them but the mandate, tools, records, and boundaries you give each one.

Chapter 2: Principles Before Machinery

It is tempting to begin with tools. Resist that temptation for one chapter, because every architectural decision in this guide flows from a small number of principles, and when the tools change, as they will within months, the principles will still be holding the structure up.

Purpose governs power. An agent should exist because a definite service requires it, not because building it was possible. Before creating any agent, write a sentence of the form: this agent exists so that a specific person or purpose is better served in a specific way. If the sentence cannot be written, the agent should not be built. An

ecosystem assembled this way stays small, legible, and governable. One assembled by enthusiasm becomes a menagerie.

Trust is earned through verification, not granted through hope. We extend trust to people gradually: first small responsibilities, observed closely; then larger ones, reviewed periodically. Extend trust to agents the same way. A new agent drafts and a human sends. Only after its judgment has been observed across many cases does it earn the right to act within defined limits, and even then its actions remain recorded and reviewable. The technical expressions of this principle are called human-in-the-loop gates, least-privilege access, and audit logs, and you will meet them in Part Three, but they are simply the machinery of earned trust.

Coherence of word and deed. An agent's conduct should be consistent with the values of the person it serves, and the ecosystem's parts should be consistent with each other. This is why your agents will share one written statement of principles rather than each carrying a private copy that drifts. It is also why this guide includes an agent whose entire role is reviewing whether plans and actions remain in harmony with stated values: coherence is not achieved once, it is maintained.

Moderation in all things, including automation. Every hour of work you delegate is an hour of judgment you must still exercise somewhere, usually in review. Automating more than you can conscientiously oversee does not multiply your capacity; it multiplies unreviewed action. The discipline of budgets, quotas, and deliberate limits, treated later under monitoring, is the practical form of this moderation.

Transparency and the written record. Institutions that endure keep minutes. An ecosystem of agents endures on the same habit: every significant action, decision, and report is written down in a place all members and the human steward can read. This principle has an architectural consequence so large that Chapter 13 is devoted to it: your agents will coordinate primarily through a shared, versioned repository of

records rather than through private conversation, because what is written can be audited, corrected, and learned from, and what is merely said between two programs vanishes.

Consultation before consequence. Actions that are hard to reverse, sending a message in your name, spending money, deleting records, signing anything, are preceded by consultation with the human steward. This is not a temporary concession to immature technology. It is a permanent feature of any arrangement in which one party acts on behalf of another.

Service to others as the test. The final measure of the ecosystem is not its sophistication but whether it enlarges your capacity to serve: your patients, your students, your family, your community. An afternoon freed from paperwork matters because of what fills it.

If you teach nothing else from this guide, teach this chapter. Learners who internalize these principles make good architectural decisions on tools they have never seen before. Learners who skip to the tools make fragile systems with whatever is fashionable.

Chapter 3: The Shape of a Personal Agent Team

Here is the destination in concrete form: a team you might actually build, and the one this guide uses as its running example. Each role answers a real need, and together they illustrate every architectural idea in Parts Two and Three.

The Counsel reviews documents with legal weight before you sign or send them: engagement letters, data-use agreements, leases, consulting contracts, publishing agreements. It reads against a checklist you and it have refined together, flags unusual terms in plain language, drafts questions to ask, and files a dated

memorandum of its review in the shared record. Two boundaries define it. First, it advises and never signs; execution belongs to you, and anything unusual earns the recommendation to consult a licensed attorney, whom the agent serves by preparing an organized brief rather than replaces. Second, because the documents it reads are confidential, it runs in the most protected part of your ecosystem, a placement Chapter 10 explains.

The Harmony Reviewer examines your plans for exercise, rest, and daily rhythm against the principles you have written down for yourself: consistency over intensity, recovery honored as training, ambition moderated by the season of life you are in. Each week it reads the record of workouts actually done, compares plan with principle and principle with practice, and writes a short reflection: where word and deed aligned, where they diverged, and what one adjustment would restore harmony. Notice what this agent is not: it is not a coach imposing a program, and it is not a medical adviser. It is a mirror with a good memory, and its value lies in the questions it raises. For readers in health care, this role generalizes: the same pattern reviews whether a clinic's schedule matches its stated commitment to unhurried visits, or whether a research group's practice matches its data-stewardship policy.

The Curator tends your learning. It watches sources you have chosen, gathers what is worth keeping about AI agents and their unfolding possibilities, summarizes each item at a stated depth, and organizes the collection so that a newcomer could walk a sequenced path through it rather than face a heap. Chapter 18's learning library is the kind of artifact a Curator maintains; once built, yours keeps it current.

The Evaluation Team watches the watchers. On a regular cadence it reviews every other agent's work: how often the Counsel's flags proved apt, what each agent cost in tokens and time, where an agent looped wastefully or overstepped its mandate, whether two agents are duplicating effort. It writes a plain-language report with recommendations: retire this agent, narrow that one's scope, move this recurring

task from an expensive model to a cheap one. In an enterprise this function is coming to be called agent operations; in a hospital you would recognize it as the quality and safety committee. Chapter 15 gives it instruments.

The Orchestrator is the attending physician of the service. It receives your requests, decides which specialist the work belongs to, ensures the shared record is read before work begins and written after it ends, notices when a task needs two agents in sequence, and brings to your attention exactly those matters that need your judgment, no more. Importantly, the Orchestrator coordinates chiefly through the shared record rather than through elaborate direct negotiation between agents, for reasons of auditability and simplicity that Chapter 14 defends.

Around these five stand two quieter structures rather than personalities. **The Registry** is the shared repository itself: the common record every agent reads and writes, the constitution of the whole. **The Gatehouse** is the set of boundaries and monitors governing what enters and leaves: which agent may use which tools, what data may cross from private to public, and the dashboard on which all activity is visible.

The clinical analogy is worth making explicit for teaching, because everything corresponds. A well-run hospital service has specialists with defined scopes of practice, an attending who coordinates, a chart that every clinician reads before acting and writes after acting, a quality committee that reviews outcomes, credentialing that limits who may do what, and infection-control practices that everyone observes not because each interaction is suspect but because the system must remain safe even when something goes wrong. Your agent ecosystem is that service in miniature, and a learner who has worked in any well-ordered institution already understands its logic.

What remains is to learn the materials. Part Two climbs the six steps of ways to build an individual agent, from a five-minute conversational bot to a private agent running entirely on your own hardware. Part Three then assembles agents into the ecosystem: the private and public tiers, the languages agents share, the firewalls between them, the common record, the orchestration, and the watching over the whole.

Part Two: The Six Steps

There are six steps from "I can describe what I want" to "I run my own agents on my own hardware." Each step is a legitimate destination, not merely a stop along the way; many people will do their best work on the second or third step and never need the sixth. For each step this part explains what it is, who it suits, what you would build with it, what it costs, and where its limits lie. A summary table closes the part.

A dated caution before the tour: this landscape changed more in the eighteen months before this writing than in the five years prior, and several flagship products were renamed, restricted, or retired during 2025 and 2026. Product specifics below were verified in late August 2026; Chapter 19 lists the items most likely to shift and should be re-checked before each teaching.

Chapter 4: Step One. Conversational Bots: Describe It and It Exists

The lowest step asks only that you write clear instructions in plain language. You are not yet building something that acts in the world; you are building a specialist that thinks with you, consistently, on demand.

Grok (xAI) is the entry point this guide recommends first, for the simple reason that its free tier is generous and its concepts transfer everywhere. Grok offers custom instructions (standing guidance up to about twelve thousand characters that shapes every conversation), and, more usefully for our purposes, Workspaces: separate environments in one account, each with its own instructions, its own uploaded files, and its own conversation history. A Workspace called "Counsel" loaded with your contract checklist and reviewed examples is a first, real version of the Counsel agent, built in an evening. Grok Automations, released in mid-2026, add scheduling: a saved prompt that runs on a daily, weekly, or monthly rhythm with results collected in a run history, free of charge, which turns a specialist that waits for you into one that reports to you. Alongside all this, xAI launched Grok Bot in beta in August 2026: always-on agents with a persistent cloud computer that can sign into your tools and carry out multi-step work, with creation as simple as naming a bot and giving it a job description. Grok Bot properly belongs on the second step, and it arrives with cautions noted there.

Claude Projects and Skills (Anthropic) are the natural next step, and the pairing matters. A Project is a workspace with standing instructions and uploaded knowledge, like a Workspace above. A Skill is something newer and, for this guide's purposes, quietly important: a folder containing a SKILL.md file (a short document of instructions, optionally accompanied by scripts and reference files) that Claude loads automatically whenever the task at hand matches the skill's description. Skills

are portable across the Claude products, from the chat interface to the command-line tools to the developer kit, and the format was published as an open standard in late 2025. The significance is this: a skill is an agent's training made into a document. When you later move your Counsel from step one to step four, its accumulated craft moves with it as files, not as memories locked in one vendor's product. Anthropic's paid plans, from roughly twenty dollars monthly, are required.

Gemini Gems (Google) offer custom instructions plus up to ten knowledge files, free of charge, with live synchronization to Google Docs and Sheets and the ability to share a Gem with colleagues as you would share a document. For a team standardizing on Google Workspace, Gems are the frictionless choice.

A note on Custom GPTs (OpenAI): for years the best-known name on this step, they were closed to new creation by individual accounts in August 2026 as OpenAI shifted this capability to business tiers under the name Workspace Agents. Existing GPTs continue to run. Teach the concept; point new individual builders to the three options above.

What this step cannot do defines the next: these are custom advisers, not custom doers. Apart from Grok's scheduled automations, they act only while you are talking to them, and they touch nothing outside the conversation.

Chapter 5: Step Two. Consumer Agent Products: The Work Gets Done

The second step is where delegation begins in earnest: products in which the agent operates a computer, works with real files, runs on schedules, and produces finished artifacts, while you review checkpoints instead of supervising keystrokes.

Claude Cowork (Anthropic), generally available since April 2026 inside the Claude desktop application, connects an agent to folders on your computer and to your connected services. Pointed at a folder of scanned invoices, it can rename them consistently, extract totals into a spreadsheet, and draft the follow-up letters, on a weekly schedule if you wish. It shares a usage quota with the rest of a paid Claude plan. For readers of this guide it holds particular interest because skills, projects, scheduled tasks, and connected folders make it the gentlest on-ramp to the ecosystem architecture of Part Three.

ChatGPT Work (OpenAI), launched July 2026 as the successor to the earlier Operator and agent-mode experiments, plans a task, asks approval, then works for as long as the task requires, drawing on a large catalog of application connections and producing documents, spreadsheets, and small applications. It is available on all tiers including free, with usage metered.

Grok Bot (xAI), in beta since August 2026, is the most ambitious framing on this step: persistent agent "teammates" that keep working while you are away and can hand tasks to one another. Two cautions are owed to learners. Pricing and availability were shifting week to week at this writing, with access tied to premium tiers and usage billed beyond them. More important, xAI's own documentation warns plainly that separate bots must not be treated as a security boundary, since they share underlying infrastructure; remember that sentence when you reach Chapter 12, because it illustrates precisely why this guide builds boundaries with architecture rather than with product features.

Perplexity Comet is a web browser with an assistant and agent built in, free at the base tier, capable of multi-step errands across websites. It earns its mention here partly as a teaching example: a publicized 2025 attack called CometJacking, in which a malicious webpage hijacked the agent through hidden instructions, is the clearest public illustration of the prompt-injection risks Chapter 12 addresses.

Who this step suits: anyone with real recurring work involving files, email, browsing, and documents, and no wish to program. Its limits: you accept each vendor's meter, model choices, and connector catalog, and the agent's memory and records live inside the product unless you deliberately externalize them, which Chapter 13 will teach you to do regardless of step.

Chapter 6: Step Three. Automation Platforms with Agent Features

The third step suits the operations-minded builder: visual platforms in which you draw the flow of work between your applications and place an intelligent step wherever judgment is needed.

n8n deserves first mention for three reasons: its agent support runs deepest, with persistent memory and some seventy AI-related building blocks; it can be self-hosted at no license cost, with unlimited runs; and it connects to local models through Ollama, making it the one step-three platform that can participate fully in the private tier of Part Three. Its cloud version begins around twenty dollars monthly. A representative build: an agent that watches a shared inbox, classifies each request, updates the patient-outreach tracker or the CRM accordingly, and alerts a human only for the exceptions.

Zapier Agents brings agents to the largest catalog of application connections, roughly nine thousand, with a free tier of several hundred activities monthly and paid tiers beyond. Its pricing unit, the activity, counts each tool use, so costs track how busy an agent is, not how useful; watch this in the Evaluation Team's reports. **Make** is the value option for classic automation, with its agent builder still maturing at this writing. **Lindy** packages the "AI employee" experience, strongest around email, meetings, and telephone. **Gumloop** and **Relay.app** serve visual pipeline building and

human-in-the-loop workflows respectively; Relay in particular has made the approval step a first-class feature, which aligns with this guide's principles.

The recurring lesson of this step is economic: credit-based and per-activity pricing rewards the discipline of Chapter 2. Agents given vague mandates wander, and wandering is billed.

Chapter 7: Step Four. The Developer's Companions

The fourth step asks basic comfort with a terminal or code editor and repays it with far more control. These tools were built for software work, but their machinery, agents with file access, scheduled tasks, and version control, serves any records-based work, and several chapters of Part Three assume a step-four tool as the ecosystem's workbench.

Claude Code (Anthropic) is a full agent that lives in the terminal (with desktop and web surfaces as well) and, for this guide, the recommended workbench. The reasons are specific. It reads and writes ordinary files in ordinary folders, which is exactly the substrate the shared record of Chapter 13 uses. It supports skills, the same portable folders introduced on step one. It supports subagents, letting an orchestrating session delegate to focused workers. It runs scheduled tasks. And it speaks the connector protocol of Chapter 11 natively. Plans from twenty dollars monthly, or metered use through the developer API.

Cursor (Anysphere) is a code editor with background agents that work in parallel while you do something else; it has become entwined with the xAI story noted in Chapter 19's volatility list. **GitHub Copilot's coding agent** accepts an assigned issue and returns a reviewed, tested pull request, a pattern worth teaching because it is pure delegation-with-review: the agent proposes, the human disposes. **Replit Agent**

builds and deploys small applications from descriptions; its effort-based pricing, in which cost is known only after the work, has produced well-publicized surprises and is a useful cautionary tale for the budgeting discipline of Chapter 15.

A concrete step-four build, achievable in an afternoon: a Claude Code scheduled task that runs each evening, reads the day's additions to your shared record, updates a summary dashboard, commits the result, and emails you three sentences. That small loop, files in, judgment applied, files out, record kept, is the whole ecosystem in miniature.

Chapter 8: Step Five. Frameworks: Writing Agents as Software

The fifth step is for building agents as durable software: programs you version, test, and run on your own schedule, with the agent loop under your control. All the major frameworks are free and open source; you pay only for the model calls they make.

The Claude Agent SDK (Anthropic) packages the same engine that powers Claude Code as a library for Python and TypeScript: file access, tool use, subagents, context management, and budget caps you can set in code. If your workbench is Claude Code, this is the natural graduation path: the prototype you shaped interactively becomes a program with the same behavior.

LangGraph (LangChain) models an agent system as an explicit graph of steps with durable state and checkpoints, the strength you want when a workflow must pause for human approval and resume days later, exactly the consultation-before-consequence pattern of Chapter 2. It is the most widely deployed of the frameworks and the most verbose to learn. **CrewAI** organizes agents as a crew of named roles with processes between them, the fastest path to a working multi-agent prototype

and a natural fit for teaching the team metaphor, though it resists unusual control flow. **The OpenAI Agents SDK** is a minimal, elegant library built around handing conversations between specialist agents with guardrails at the boundaries.

Microsoft's Agent Framework reached general availability in April 2026 as the merger of its two predecessors, AutoGen and Semantic Kernel, both now in maintenance; teach the merged framework, not the ancestors. **Google's Agent Development Kit (ADK)** is code-first with the broadest language support and first-class support for the agent-to-agent protocol of Chapter 11. **smolagents (Hugging Face)** is a few hundred lines of framework in which the agent writes small programs as its actions; it pairs naturally with local models and demands proper sandboxing, making it excellent teaching material for Chapter 12.

Choosing among them matters less than learners fear. The concepts, mandates, tools, loops, checkpoints, records, transfer completely; the shared record of Chapter 13 is framework-neutral by design; and the protocols of Chapter 11 exist precisely so that a CrewAI agent and a LangGraph agent can serve one another. Choose by fit: LangGraph for durable stateful processes, CrewAI for quick team prototypes, the Claude Agent SDK for deep integration with a Claude-based workbench, ADK for a Google-centered stack.

Chapter 9: Step Six. Private Agents on Your Own Hardware

The final step runs the mind itself, the language model, on hardware you control, so that certain work never leaves your desk. For a physician holding patient records, a lawyer holding privileged documents, or anyone who simply believes some matters are nobody else's business, this step is not an enthusiasm; it is the load-bearing wall of the private tier in Chapter 10.

The runtimes are settled and free. **Ollama** is the developer's default: one command downloads and serves any of a hundred-plus prepared models behind a standard local interface that every framework on step five can call. **LM Studio** wraps the same capability in a friendly graphical application, the right first stop for the non-programmer. **llama.cpp** underlies much of the ecosystem and rewards those who want maximum control.

The open models worth teaching in late August 2026: **Qwen3-Coder 30B** as the best balance of capability to hardware demands, running well on a strong consumer graphics card or a 32-gigabyte Apple machine; **gpt-oss** (OpenAI's open-weights family) at 20 billion parameters for 16-gigabyte machines; **Devstral 24B** with genuinely published agentic benchmarks; and the **DeepSeek** family for reasoning at low cost. A correction for older curricula: Meta's Llama line, long the default recommendation, ceased to be so after Meta's 2026 turn toward closed models; do not build a fresh teaching on it.

Hardware, in one paragraph. A model's appetite is measured in gigabytes of memory on the graphics card or unified memory, and quantization (storing the model's numbers more coarsely, at four or eight bits rather than sixteen) roughly halves or quarters the appetite at a modest cost in quality. Eight gigabytes runs small 7-billion-parameter assistants; sixteen runs the 20-billion class; twenty-four to thirty-two, a used high-end graphics card or a well-equipped Mac, is the current sweet spot for capable local agents; and a machine without a graphics card, running on its main processor alone, produces a handful of tokens per second, too slow for agent loops. Budget six hundred to two thousand dollars for a dedicated machine, or repurpose a recent Mac.

Two assemblies show the step's range. Self-hosted n8n plus Ollama yields a fully private, no-code agent platform at zero license cost. **OpenClaw**, the breakout open-source personal agent of 2026, connects any model, including local ones, to your

messaging channels and computer, with persistent memory and an extensible skill system; its explosive popularity arrived together with sobering security incidents from carelessly exposed installations, and it should be taught with Chapter 12 in the same breath.

The six steps at a glance

Step	You provide	Representative tools	Typical cost	Ceiling
1. Conversational bots	Clear written instructions	Grok Workspaces and Automations, Claude Projects and Skills, Gemini Gems	Free to ~\$20/mo	Advises; does not act
2. Consumer agent products	Real tasks and review time	Claude Cowork, ChatGPT Work, Grok Bot, Comet	\$0 to \$200/mo, often metered	Vendor's meters, models, and connectors
3. Automation platforms	Process thinking	n8n, Zapier Agents, Make, Lindy, Relay	Free tiers; ~\$9 to \$50/mo	Bound to connector catalogs; per-activity billing
4. Developer companions	Terminal comfort	Claude Code, Cursor, Copilot coding agent, Replit	~\$20 to \$200/mo or metered	Software-shaped workflows
5. Frameworks	Programming	Claude Agent SDK, LangGraph, CrewAI, OpenAI Agents SDK, MS Agent Framework, Google ADK, smolagents	Free; pay per model call	Your engineering time
6. Local and private	Hardware and care	Ollama, LM Studio, Qwen3, gpt-oss, DeepSeek, OpenClaw, n8n self-hosted	Hardware \$600 to \$2,000; then ~\$0	Model quality below frontier; your own operations

Read the table vertically one more time and notice what increases as you go further: not capability alone, but responsibility. On step one the vendor holds nearly every risk; on step six you hold nearly all of it. The architecture of Part Three exists so that you can hold it well.

Part Three: The Architecture of an Ecosystem

A single agent is a tool. An ecosystem is an arrangement: private and public parts in right relation, a common language, boundaries that hold even when something misbehaves, a shared record that makes the whole legible, coordination that stays simple, and instruments that let the steward see. This part takes those six in turn.

Chapter 10: Private and Public: Two Realms in Right Relation

The first architectural decision is not which model to use but where each kind of work is allowed to happen. Picture your ecosystem as a household with three concentric spaces.

The innermost space is **your own machine**: local models on your hardware, files on your disk. Work done here generates no bill, travels over no network, and answers to no one's terms of service. Its models are capable but not the most capable. This is where the Counsel reads contracts and where anything touching health records, finances, or the confidences of others belongs by default.

The middle space is **your accounts in the cloud**: the paid Claude, Grok, or OpenAI services operating under your credentials and their privacy terms. Most everyday agent work lives here, where frontier-quality judgment is available on tap and the pricing is metered.

The outermost space is **the public web**: search, fetched pages, third-party services, other people's agents. Information from this space is valuable and, as Chapter 12 will insist, permanently untrusted in the technical sense: it may contain instructions that are not yours.

The craft is routing: deciding, for each piece of work, the innermost space that can do it well. Practitioners route along three axes. **Sensitivity** rules first: a request touching protected information is answered locally or not at all, and this rule must fail closed, meaning that if the local model is down, the request fails with an error rather than quietly escaping to the cloud. It is worth pausing on how easily the opposite happens by default: most software is written to fall back, to retry elsewhere, to be helpful; a privacy boundary must be written to refuse. **Complexity** rules second: try the small local model first for classification, extraction, formatting, and routing, and escalate to a frontier model only when confidence is low. Field reports through 2026 consistently find that most requests inside agent systems never needed a frontier model, which is why hybrid arrangements commonly cut costs by more than half. **Latency** rules third, since a local model answers in milliseconds and a cloud round trip takes seconds; loops that run hundreds of times a day belong near the metal.

Concretely, the hybrid is assembled from parts you have already met: Ollama serving local models; a small routing layer, either a gateway such as LiteLLM or simply your framework's branching logic, holding the routing rules in one auditable place; a screening step that scans outbound requests for identifiers (names, record numbers, account numbers; Microsoft's open-source Presidio is the standard instrument) before anything crosses to the middle space; and a log of every routing decision. For

a health-care teaching audience the summary is one sentence: the architecture enforces the minimum-necessary principle you already practice, with the local machine as the covered entity and the cloud as the business associate you engage deliberately, under terms, for defined purposes.

Chapter 11: A Common Language: How Agents Connect

For agents from different makers to work together, they need shared conventions, and by 2026 the field has largely settled on two, both now stewarded by a neutral body, the Agentic AI Foundation under the Linux Foundation, to which Anthropic donated the first and Google the second.

The Model Context Protocol (MCP) standardizes how an agent connects to tools and data. The analogy that lands with every audience is the universal connector: before MCP, connecting three assistants to five services meant fifteen bespoke integrations; after it, each service publishes one MCP server and every assistant that speaks the protocol can use it. Health-care readers may prefer a closer analogy: MCP is to agent-tool connections what FHIR is to health-data exchange, a common interface that ends the era of custom point-to-point plumbing. An MCP server exposes tools (actions the model may invoke), resources (data it may read), and prompts (templates a user may call). Servers run in two ways, and the difference matters architecturally: a local server runs as a child process on your machine, meaning the tool and its data stay inside your inner space; a remote server is reached over the web. The protocol's adoption is the reason cross-platform interoperability exists in practice: Claude, ChatGPT, Grok, Gemini, Copilot, Cursor, and virtually every serious agent product is an MCP client today. Grok added bring-your-own-MCP

connectors in May 2026; ChatGPT's developer mode grants full MCP access with the vendor's own blunt warning that write-capable tools are powerful and dangerous.

The Agent2Agent protocol (A2A) standardizes how one agent engages another as a peer, without either exposing its inner workings. Each agent publishes an Agent Card, a small signed document at a well-known address describing who it is, what skills it offers, and how to authenticate; work between agents is structured as tasks with a defined lifecycle (working, input required, completed, failed), and the states are worth noticing because input-required is consultation-before-consequence built into the wire format. A2A reached its first stable version in March 2026 and joined MCP under the same foundation in August 2026. Its honest status: the standard for framework-to-framework and enterprise use, demonstrated across the major step-five frameworks, while consumer products mostly do not yet speak it to each other.

Two humbler conventions complete the picture. **AGENTS.md**, now used in tens of thousands of repositories and stewarded by the same foundation, is simply a markdown file of instructions that any visiting agent reads first, a doormat that says how we do things here; Anthropic's **CLAUDE.md** convention is its ancestor, and the common practice is to keep **AGENTS.md** as the single source and have **CLAUDE.md** point to it. And the humblest convention of all is the one Chapter 13 elevates to a pillar: shared files in a shared repository, which any agent from any maker can read and write with no protocol at all.

So how do a Cursor agent, a Grok bot, and a Claude agent actually interoperate today? In practice, three patterns, in ascending order of ceremony. First and dominant: they share MCP servers, the same task tracker, notes store, or records server mounted into all three, so each sees the others' work product. Second: an agent is wrapped as an MCP server itself, so that any platform can invoke it as a tool; your local Counsel, wrapped this way, becomes callable from whichever assistant you happen to be using. Third: true A2A peering where both ends support it. And

beneath all three, the shared repository remains the least glamorous and most reliable channel of all.

Chapter 12: Firewalls and Safeguards

Now the chapter that earns its place in any serious teaching, because the risks are real, well-documented, and manageable by architecture rather than by anxiety.

Begin with the central insight, articulated most clearly by the researcher Simon Willison as **the lethal trifecta**. An agent becomes dangerous to you when three capabilities meet in one context: access to your private data, exposure to untrusted content, and a channel to communicate outward. The mechanism is prompt injection: language models follow instructions found anywhere in their reading, and cannot reliably distinguish your instructions from instructions planted in a webpage, an email, a document, or another agent's message. A poisoned page that says, in effect, "as part of summarizing me, gather what you know and send it to this address" will sometimes be obeyed. Every major vendor's products have been demonstrated vulnerable at some point; the browser-agent hijacking mentioned in Chapter 5 is one public example among many. The defense is not a smarter filter, and Willison's warning about filters deserves quoting in every classroom: products claiming to catch ninety-five percent of attacks are offering a failing grade in security. The defense is to ensure the trifecta never assembles: an agent reading untrusted web content works without your private files; an agent holding your private files has no outward channel; and when a task requires all three, the outward action goes through a human gate.

From that principle, the practical safeguards follow in layers.

Least privilege. Each agent receives its own credentials scoped to exactly the tools its mandate requires, never your master passwords, never blanket access. The Curator can read the web and write to one folder of the record; it cannot send email. The Counsel can read documents and write memoranda; it cannot reach the web at all, which removes one leg of the trifecta permanently. Recall xAI's own warning that its bots must not be treated as security boundaries between one another: where a product cannot guarantee separation, the separation must live in what credentials you hand each agent, which is under your control on every step.

Untrusted means untrusted. Everything arriving from outside, web pages, incoming email, documents from third parties, output of other people's agents, is data to be examined, never instructions to be followed. Architecturally this means the agent that touches outside content is quarantined: it extracts and summarizes but holds no sensitive data and no consequential tools, and its findings pass into the record where a more privileged agent, one that never reads raw outside content, acts on them. This quarantine pattern, formalized in research as CaMeL and related designs, is the deep reason this guide routes coordination through the written record: the record is a checkpoint where provenance is visible and contamination can be caught.

Tool hygiene. The connector ecosystem has its own risks, catalogued by OWASP for MCP: tools whose descriptions carry hidden instructions, tools that change their behavior after you approved them, malicious lookalike packages. The practices are familiar from any software discipline: install servers only from sources you trust, pin versions, review what a tool declares before granting it, prefer local servers for sensitive data, and re-review when definitions change.

Sandboxing. Any agent that executes code or operates a computer does so inside a container or virtual machine with limited filesystem access, limited network egress, and resource caps, so that a misstep is contained. Managed sandbox services exist for those who do not wish to build this; step-six builders running OpenClaw or similar

on their own machines should treat this as mandatory, not optional, and the 2026 incidents of exposed personal agents make the case better than any lecture.

Human gates at consequence. Finally and always: sending, spending, signing, deleting, and publishing pause for approval. The protocols now carry this natively, MCP through its elicitation mechanism, A2A through the input-required state, and every serious product exposes an approval mode. Configure it, and resist the temptation to approve reflexively; the gate is where your judgment enters the system, and Chapter 15's evaluation reports will tell you if approvals have become rubber stamps.

Taught in one line: give each agent the least it needs, treat the outside world as data, keep consequences behind a human gate, and let the boundaries be enforced by architecture, not by trust in any single product or filter.

Chapter 13: The Common Record: A Repository as Shared Memory

Here is the guide's central architectural commitment, the one that makes everything else simpler: your agents coordinate primarily through a shared, versioned repository of plain files, hosted on GitHub, rather than through direct conversation with each other.

The pattern is old and honorable. Computer science calls it a blackboard architecture: specialists gather around a common board, each reads the current state of the problem, contributes what its expertise allows, and writes the contribution back for all to see. Medicine has practiced it for a century under another name: the chart. No competent clinician acts on a patient without reading the chart first, and none finishes without writing in it. The chart, not the hallway conversation,

is the coordination mechanism of record, and for the same reasons your agents' repository will be.

Why versioned files rather than messages between agents? Five reasons, each of which is a principle from Chapter 2 wearing working clothes. **Auditability:** every change to a repository is a commit bearing author, timestamp, and an exact difference; the question "which agent changed this recommendation, when, and from what" is answered by the version history in seconds, while a message exchanged directly between two agents is gone. **Durability:** agents forget everything between sessions; the desk is cleared, but the record remains, so any agent can resume any thread by reading. **Simplicity of joining:** adding a sixth agent to a message-passing web means wiring it to five others; adding it to the record means telling it where the record is. **Visible conflict:** when two agents disagree about a file, version control surfaces a merge conflict for resolution, instead of one silently overwriting the other. **Reviewability by the steward:** you can read your agents' working record exactly as you read a colleague's, and the Evaluation Team's raw material is simply the history. The field converged on this pattern through 2025 and 2026: production agent-memory systems rebuilt themselves on git, coding agents adopted repository-based issue trackers as shared memory, and the influential engineering essays of the period argue for single-writer, file-based state over agent-to-agent chatter. Direct messaging between agents still has a place, and Chapter 14 gives it one, but it is the exception, justified case by case, never the default.

A structure that serves well, and which you should adapt rather than adopt:

```

agent-registry/
├─ AGENTS.md                # The constitution: principles, conventions,
│                            #   and standing instructions every agent reads first
├─ CLAUDE.md                # One line: @AGENTS.md
├─ charters/                # One file per agent: mandate, tools granted,
│   ├── counsel.md          #   boundaries, escalation rules
│   ├── harmony.md
│   ├── curator.md
│   ├── evaluator.md
│   └── orchestrator.md
├─ state/
│   └── STATE.md             # Current status and handoffs, ~100 lines,
│                           #   rewritten at the end of every session
├─ records/                 # Append-only dated working records
│   ├── 2026-08-31-counsel-nda-review.md
│   └── ...
├─ decisions/               # Numbered records of decisions made, immutable
│   └── 0007-local-only-for-legal-documents.md
├─ reports/                 # Generated outputs bound for the steward
│   ├── weekly-harmony.md
│   └── monthly-evaluation.md
└─ .github/workflows/       # Automation, including the Drive push below

```

Three disciplines make the structure live. Every agent session begins by reading AGENTS.md, its own charter, and STATE.md, and ends by writing its record and refreshing STATE.md; write the contract into every charter. Records are dated, appended, and never silently edited; corrections are new entries. And decisions, once accepted, are immutable, exactly like the architecture decision records of software practice: the reasoning of the past remains legible even when superseded. When multiple agents work simultaneously, git worktrees, a standard feature that gives each agent its own working copy of the same repository, prevent them from treading on one another; the technique is now documented practice in every major coding-agent tool.

From the record to the steward: pushing reports into Google Drive. A repository is the agents' home ground, but you may live in Drive, and reports should come to

where you are. Four routes, in ascending robustness. A Zapier or Make template copies new repository files to Drive in minutes with no code, at the price of task limits and shallow file handling. A community GitHub Action (adityak74/google-drive-upload-git-action is the most used) uploads on every push with one short workflow file; its known trap is that Google service accounts own no storage of their own, so on a personal account you must target a shared drive or overwrite an existing user-owned file. A Google Apps Script on a timer, pulling from the GitHub API and writing with your own identity, is the best zero-infrastructure answer for a personal Google account. And the most robust: rclone, the veteran synchronization tool, run inside a GitHub Actions workflow with your own authorized configuration stored as a secret, giving idempotent one-line synchronization of the reports folder to a Drive folder. A worked example appears in Chapter 16.

Chapter 14: The Orchestrator: Coordination Without Committees

With a common record in place, orchestration becomes lighter than most writing on multi-agent systems suggests, and the first duty of this chapter is to pass on the field's hardest-won lesson: use fewer agents than you think.

Two essays frame the debate for any classroom. Anthropic's account of its multi-agent research system shows the pattern working: an orchestrator decomposes a hard question, dispatches parallel workers, and synthesizes, at roughly fifteen times the token cost of a single conversation. Cognition's rejoinder, "Don't Build Multi-Agents," shows the failure mode: agents working from fragments of context, contradicting one another, and compounding errors, and argues for one agent with one continuous context wherever possible. The synthesis this guide teaches: parallelize only work that is truly separable and read-heavy (research, review,

gathering); keep work that builds one coherent artifact under one agent's hand; and when in doubt, one careful agent with a good record beats three coordinating ones.

The Orchestrator's actual duties, then, are modest and vital. It receives your intentions and translates them into task entries in the record. It routes each task to the right specialist, using the routing axes of Chapter 10. It sequences: the Counsel's review precedes the signature; the Curator's gathering precedes the weekly digest. It enforces the session contract, record read before work, record written after. It escalates to you exactly what its charter says must be escalated. And it convenes the rare true collaboration, standing sequences the field has named: prompt chaining (output of one specialist feeds the next), routing (classify, then dispatch), orchestrator-workers (decompose, parallelize, synthesize), and evaluator-optimizer (one agent drafts, another critiques, the first revises; your Evaluation Team runs this loop on the whole ecosystem monthly).

How you embody the Orchestrator depends on your step, and every step can have one. On steps one and two it is you, aided by a coordination workspace whose instructions encode the routing rules; Grok Bot's handoffs and Cowork's scheduled tasks automate pieces. On step four it is a Claude Code session with the orchestrator charter loaded, dispatching subagents. On step five it is a LangGraph supervisor graph or CrewAI hierarchical process. In every embodiment, hold the same line: the Orchestrator coordinates through the record, keeps no private state the record does not hold, and is itself subject to the Evaluation Team's review like any other member.

Chapter 15: Watching Over the Whole: Monitoring and Evaluation

An ecosystem you cannot see is one you cannot govern. This chapter equips the steward, and the Evaluation Team, with instruments: first observation,

then judgment.

Observation. The unit of observation is the trace: the complete recorded path of one task through the system, every model call, every tool invocation, every token counted and priced, every intermediate decision, arranged as a timeline you can replay. A trace is to an agent what the flight recorder is to an aviation review or the operative note is to a morbidity conference: not surveillance, but the precondition of learning from practice. The good news of 2026 is that observation has a common gauge: the OpenTelemetry conventions for generative AI, an open standard by which agents emit their traces in one format regardless of maker. Claude Code, Copilot, and Codex already emit it natively; the step-five frameworks all instrument to it. The practical recipe for the single dashboard the steward wants, covering agents across Claude, Cursor, Grok, and the rest: enable native telemetry where offered, instrument custom agents with the standard, run one collector as the funnel, and point it at one backend. For that backend, the individual builder is well served free: Langfuse (open source, self-hostable, generous hosted free tier) is this guide's default recommendation; Arize Phoenix (open source, uncapped self-hosting, telemetry-native) is the strongest pure-standard choice; LangSmith is natural if your stack is LangGraph; AgentOps specializes in replaying multi-agent sessions; Helicone is the five-minute proxy option. Whichever you choose, the dashboard should answer at a glance: what ran today, what did it cost, what failed, what waits for my approval, and what did any agent do that it has never done before.

Judgment. Observation tells you what happened; evaluation asks whether it was good. The field's consensus, and the Evaluation Team's method, examines four dimensions: the outcome (was the task completed correctly, judged against a reference or by a carefully instructed model acting as judge, calibrated periodically against your own judgment); the trajectory (was the path sensible, or did the agent loop, retry, and wander; wasted steps are wasted tokens and a signature of a

confused mandate); tool use (right tool, right arguments, checkable mechanically); and cost and latency against the task's worth. The working practice is the regression set: a curated collection of real past tasks with known good outcomes, re-run whenever a model, prompt, or charter changes, so that improvement somewhere is not purchased with silent regression elsewhere. Every production failure that surprises you earns a place in the set. Do not let learners over-index on public benchmarks; by 2026 the serious practice is evaluation on your own tasks, your own data, your own standards.

The stewardship function. In enterprises this is crystallizing into a named discipline, agent operations, with registries of agents, identities and scoped credentials for each, budget enforcement at gateways, on-call rotations, and kill switches. Your ecosystem needs the same functions at household scale, and they fit on one page: a monthly evaluation report (quality, cost, incidents, recommendations) written by the Evaluation Team and read by you; per-agent budgets with alerts, because an agent in a loop is a meter running; scoped credentials per agent as Chapter 12 required; a kill switch you have actually rehearsed, even if it is only revoking a key and stopping a schedule; and a standing rule that any new agent, tool grant, or schedule is itself recorded in the registry. This is governance sized to a person: not bureaucracy, but the same habit of periodic, written, principled review that keeps any institution worthy of the trust placed in it.

Part Four: Practice

Chapter 16: A Worked Example: Assembling the Team

Theory earns its keep in assembly. Here is one concrete embodiment of the five-agent team from Chapter 3, chosen for teachability; every element can be substituted, and the charters and record travel intact if you later change tools. It assumes the middle of the six steps: a paid Claude plan with Claude Code or Cowork as the workbench, a free GitHub account, a Google account, and optionally one machine running Ollama for the private tier.

The Registry comes first, because everything else reads from it. Create a private GitHub repository named agent-registry with the structure from Chapter 13. Write AGENTS.md before writing any agent: your principles from Chapter 2 in your own words, the session contract, the naming conventions for records, and the escalation rules. This document is the most important hour of the whole build, and it is an hour of writing, not programming.

The Counsel runs in the private tier. Install Ollama and pull a capable local model (qwen3-coder or a peer from Chapter 9). Its charter grants it: read access to a documents folder, write access to records/ in the registry, no network. Whether you drive it through a local-model session in a step-four tool or through a small step-five script, the working pattern is identical: you place a contract in the folder, the Counsel reads its charter and checklist, writes a dated review memorandum into records/, and the memorandum ends with the standing line its charter requires: this review assists your judgment and does not replace licensed counsel. If a fully local setup is a step too far today, an honest interim is a Claude Project with the same

charter, accepting the middle tier's confidentiality terms deliberately, and noting in the decision record that you did so and why. That is the decisions/ folder doing its work: even compromises are made in writing.

The Harmony Reviewer is a scheduled task on the workbench. Its charter names the sources: the principles file, and wherever your workout record lives (an exported file dropped weekly into the repository is the simplest reliable feed). Each Sunday it runs, compares practice against principle, and writes reports/weekly-harmony.md: what aligned, what diverged, one adjustment proposed. Keep its mandate reflective rather than prescriptive; the charter line that does this is a single sentence: you raise questions and observe patterns; you do not prescribe programs or give medical advice.

The Curator is a second scheduled task with the opposite privileges: web access, no private data, in keeping with the trifecta discipline. Weekly, it reviews the sources named in its charter, writes structured entries into records/curation/, and refreshes a reading-path document in reports/. Seed its charter with the library in Chapter 18, and add the standing instruction that every entry carries source, date, difficulty, and a two-sentence judgment of why it earned inclusion.

The Evaluation Team is a monthly scheduled task, and deliberately the most senior charter. It reads everything: the month's records, the observability export (Langfuse's data can be pulled by API or file), the commit history. It writes reports/monthly-evaluation.md: per-agent cost, task counts, failures and loops, drift from charter, duplication between agents, and up to three recommendations, each phrased as a proposed change to a specific charter file. You accept a recommendation by committing the charter change; the commit is the decision record.

The Orchestrator, in this embodiment, is your daily working session with the orchestrator charter loaded: it reads STATE.md at open, routes work to the right specialist's charter and tier, and rewrites STATE.md at close. Automate its routine dispatches gradually; keep its escalations to you sacred.

The bridge to Drive is one workflow file, using the robust route from Chapter 13. Once, locally, run rclone config to authorize your Google account, then store the resulting configuration as a repository secret named RCLONE_CONF. The workflow:

YAML

```
# .github/workflows/reports-to-drive.yml
name: Push reports to Google Drive
on:
  push:
    paths: ["reports/**"]
    workflow_dispatch:
jobs:
  sync:
    runs-on: ubuntu-latest
    steps:
      - uses: actions/checkout@v4
      - name: Install rclone
        run: curl https://rclone.org/install.sh | sudo bash
      - name: Write rclone config
        run: |
          mkdir -p ~/.config/rclone
          echo "${{ secrets.RCLONE_CONF }}" > ~/.config/rclone/rclone.conf
      - name: Sync reports
        run: rclone sync reports/ gdrive:AgentReports --create-empty-src-dirs
```

From then on, any report an agent commits appears in the AgentReports folder of your Drive within a minute or two, and the Drive folder can be shared with anyone who should see the reports without touching the repository. The pattern generalizes: the repository is the system of record, and Drive, email, or a dashboard are projections of it for human eyes.

The dashboard completes the assembly: a Langfuse account (or self-hosted instance beside Ollama for the fully private version), telemetry enabled on the workbench, and a browser tab that answers the steward's five questions from Chapter 15. Total recurring cost of the whole embodiment, beyond model subscriptions you likely already hold: at or near zero.

Build it in this order, one week per element if you are teaching it as a course: registry, one specialist end to end (the Curator is the gentlest start), the Drive bridge, the dashboard, the remaining specialists, and the Evaluation Team last, once there is a month of record for it to evaluate.

Chapter 17: A Thirty-Day Path for Learners

For students who want a sequence rather than a reference, here is a month, in four movements, requiring no programming until the final week and even then only optionally.

Days one through seven: one faithful adviser. Read Chapters 1 through 3. Write your principles page, one page, in your own hand. Then build one step-one specialist in a free tool (a Grok Workspace or a Gemini Gem): a real mandate from your actual work, a charter written before the first conversation. Use it daily for a week, refining the charter each time it disappoints you. The lesson of the week is that the charter, not the model, is where quality lives.

Days eight through fourteen: the record. Create the registry repository (GitHub's web interface suffices; no code). Move your charter into it. Begin the discipline: each significant session with your adviser ends with a dated record entry, pasted in by hand if need be. By the end of the week you will have felt the change the record makes: continuity across sessions, and the ability to see your own system.

Days fifteen through twenty-one: delegation with a gate. Add a second agent on step two or three: a scheduled or automated worker (a Grok Automation, a Cowork scheduled task, or an n8n flow) whose output lands in the record. Configure the approval gate on anything that leaves the system, and exercise it consciously all week. Read Chapter 12 mid-week, then audit your two agents against the trifecta.

Days twenty-two through thirty: the whole in miniature. Add the bridge that puts a weekly report where you and others will see it (the Zapier route needs no code; the rclone route if you are able). Then sit as your own evaluation team: read the month's record and write the first monthly report yourself, by hand, answering the four questions of Chapter 15. Deciding what the Evaluation Team should ask, by having been it, is the best possible preparation for delegating that role. Close the month by writing the charter for your next agent, and notice that you now do this fluently.

Chapter 18: The Learning Library

A curated shelf, current at this writing, ordered roughly by the path a learner would walk. Each item earned its place; the Curator's standing task is to keep this chapter alive.

Foundational documents, read in this order:

Work	Source	Why it matters
Building Effective Agents	Anthropic Engineering (anthropic.com/engineering/building-effective-agents)	The canonical starting essay: workflows vs. agents and the five composition patterns
A Practical Guide to Building Agents	OpenAI (cdn.openai.com , business guides)	The complementary vendor view, beginner-friendly
ReAct: Synergizing Reasoning and Acting	Yao et al., ICLR 2023 (arxiv.org/abs/2210.03629)	The research root of the agent loop: reason, act, observe
How We Built Our Multi-Agent Research System	Anthropic Engineering	Orchestrator-workers in production, including the fifteen-fold token finding
Don't Build Multi-Agents	Cognition (cognition.com/blog)	The essential counterpoint; pair with the previous for the design debate
Effective Context Engineering for AI Agents	Anthropic Engineering	The context-as-scarce-resource framing that now governs practice
The Lethal Trifecta	Simon Willison (simonwillison.net)	The security essay; required before anyone grants an agent tools
Why Do Multi-Agent LLM Systems Fail?	Cemri et al. (arxiv.org/abs/2503.13657)	The empirical taxonomy of fourteen failure modes
12-Factor Agents	HumanLayer (github.com/humanlayer/12-factor-agents)	Engineering principles for production-grade agents
Agents Companion whitepaper	Google (kaggle.com/whitepaper-agent-companion)	Production operations, agentic retrieval, evaluation

Structured courses, all free: the Hugging Face Agents Course

(huggingface.co/learn/agents-course) is the best complete curriculum with a certificate; DeepLearning.AI's Agentic AI (Andrew Ng's flagship on the four agentic

patterns) plus its short courses on multi-agent systems, evaluating agents, and agentic memory; Anthropic Academy (anthropic.com/learn) for the Claude API, MCP, and skills; Google and Kaggle's five-day AI Agents Intensive; and Microsoft's AI Agents for Beginners (github.com/microsoft/ai-agents-for-beginners), now expanded to eighteen lessons.

Voices worth following: Simon Willison's weblog for daily, rigorous notes; Latent Space for the profession's conversation with itself; Ethan Mollick's One Useful Thing for the managerial and educational lens; the Anthropic engineering blog for the highest-signal vendor writing; and on YouTube, Andrej Karpathy and 3Blue1Brown for foundations, IBM Technology for short concept explainers, and builders such as Cole Medin, Sam Witteveen, and IndyDevDan for working sessions.

Standards and communities: the protocol homes (modelcontextprotocol.io, a2a-protocol.org, agents.md) reward direct reading; [r/AI_Agents](https://r/ai_agents) is the largest open forum; the AI Engineer conference posts its talks freely; and the Hugging Face course Discord offers study companionship.

Chapter 19: Keeping the Teaching Alive

A document about this field is a living document or it is a misleading one. Three lists close the guide.

Re-verify before every teaching, because these were in motion at this writing: Grok Bot's pricing, availability, and architecture (weeks old and shifting, amid reported consolidation among xAI, Cursor's maker, and related companies); the successor arrangements for OpenAI's Custom GPTs and agent products, which changed names twice in eighteen months; the current open-model recommendations of

Chapter 9, where a season is a generation; per-seat and per-activity prices on every step; and the maturity of A2A adoption beyond the enterprise frameworks.

Treat as stable, because these have institutional weight behind them: MCP as the tool-connection standard and its stewardship under the Agentic AI Foundation; the repository-as-record pattern and the AGENTS.md convention; the telemetry conventions for observation; the security analysis of Chapter 12, whose logic does not depend on any product; and every word of Chapter 2.

Teach in this order, whatever the audience's level: principles, then the team as a story, then one build at the audience's step, then the record, then safeguards, then the rest as appetite demands. Learners who leave with a written charter, a repository, and the habit of the record have the whole of it; the tools they happen to use this year are costume.

Closing: The Measure of the Work

It is fitting to end where Chapter 2 began. The worth of an ecosystem of agents is not measured by its ingenuity but by what it releases you to do. Hours recovered from paperwork are the raw material; what they become, deeper attention to patients and students, work of greater scope undertaken with a clearer mind, presence restored to family and community, is the actual product. Build the smallest system that serves, govern it in writing, review it in consultation, and let its quiet order be one more expression of a life arranged around service. Then teach it, because capability of this kind, widely and wisely shared, is how a community, and not merely an individual, is enriched.

Appendix: Glossary

Agent

A language model given tools, memory, and a mandate, able to pursue a goal across multiple steps of its own choosing.

A2A (Agent2Agent)

The open protocol by which agents engage one another as peers, via published Agent Cards and structured tasks.

Blackboard architecture

Coordination through a shared record all specialists read and write, rather than through direct messages.

Charter

The written mandate of one agent: purpose, tools granted, boundaries, escalation rules.

Context window

The bounded amount of text a model holds in view at once; the desk, not the library.

Fail closed

The design rule that a privacy or safety boundary, when its preferred path is unavailable, refuses rather than falls back.

Human-in-the-loop gate

A required pause for human approval before a consequential action.

Least privilege

Granting each agent only the access its mandate requires.

Lethal trifecta

The dangerous meeting, in one agent context, of private data, untrusted content, and an outward channel.

LLM-as-judge

Using a carefully instructed model to grade another model's output, calibrated against human judgment.

MCP (Model Context Protocol)

The open standard by which agents connect to tools and data sources.

Orchestrator

The agent (or person) that routes work, sequences specialists, and enforces the session contract.

Prompt injection

The planting of instructions in content an agent reads, in hope the agent obeys them.

Quantization

Storing a model's numbers more coarsely to shrink its memory appetite at modest cost in quality.

Regression set

A curated collection of past tasks with known good outcomes, re-run after any change.

Token

The unit in which models read, write, and are priced; roughly three quarters of an English word.

Trace

The complete recorded path of one task: every call, tool use, token, and decision, replayable.

Workflow

A sequence whose steps are fixed in advance, with the model filling in the steps; contrast with agent.

This assists your judgment. It does not replace licensed professional judgment. It is architecture, not a certification.